

Fast and Expressive Multi-Byte Prediction with Probabilistic Circuits



THE UNIVERSITY
of EDINBURGH



NatWest Group

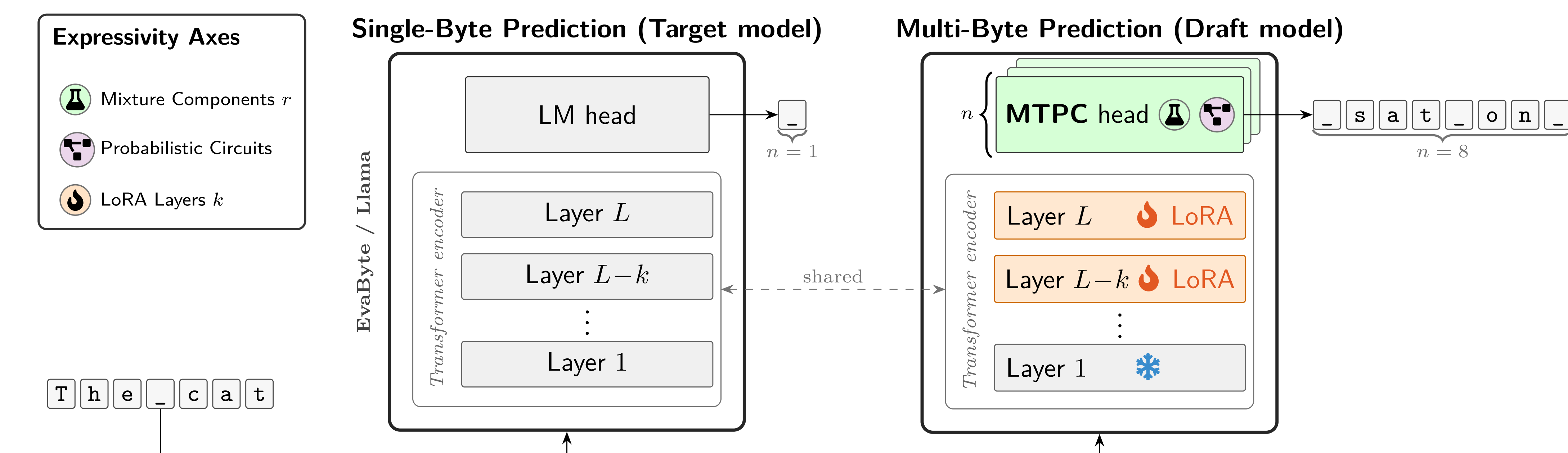


Andreas Grivas¹, Lorenzo Loconte¹, Emile van Krieken², Piotr Nawrot¹, Yu Zhao¹, Euan Wielewski³, Pasquale Minervini^{1,4}, Edoardo Ponti¹, Antonio Vergari¹
School of Informatics, University of Edinburgh¹ Vrije Universiteit Amsterdam² NatWest Group³ Miniml.AI⁴

1 Multi-Byte Prediction

Goal Faster text generation...

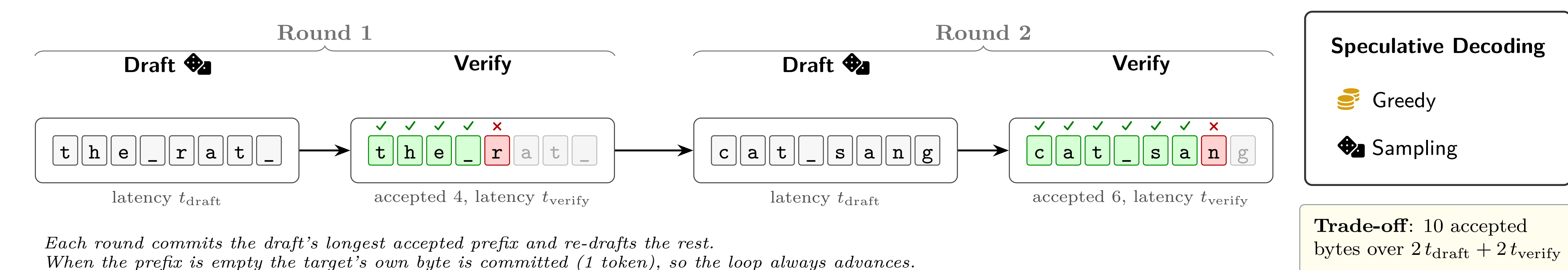
Predict n tokens at each timestep by modelling $q_{\theta}(x_{t+1}, \dots, x_{t+n} \mid \mathbf{x}_{\leq t})$.



2 Speculative Decoding

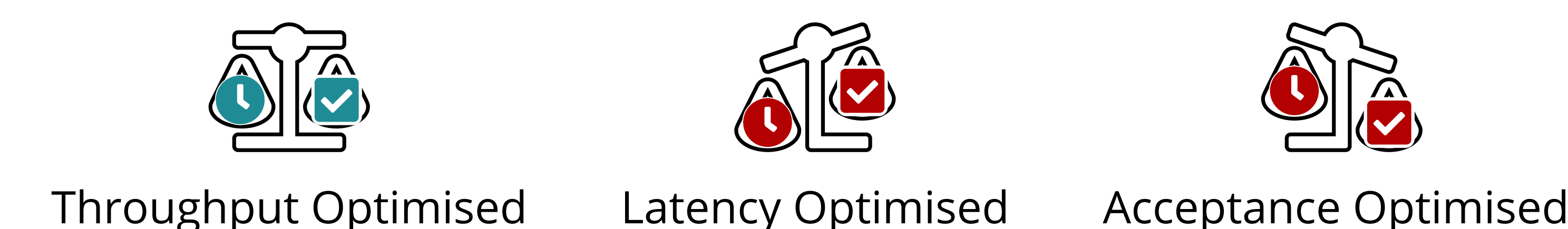
Goal ...with no loss in generation quality.

Reject drafted tokens that do not match the target model (exactly for greedy and in expectation for sampling).



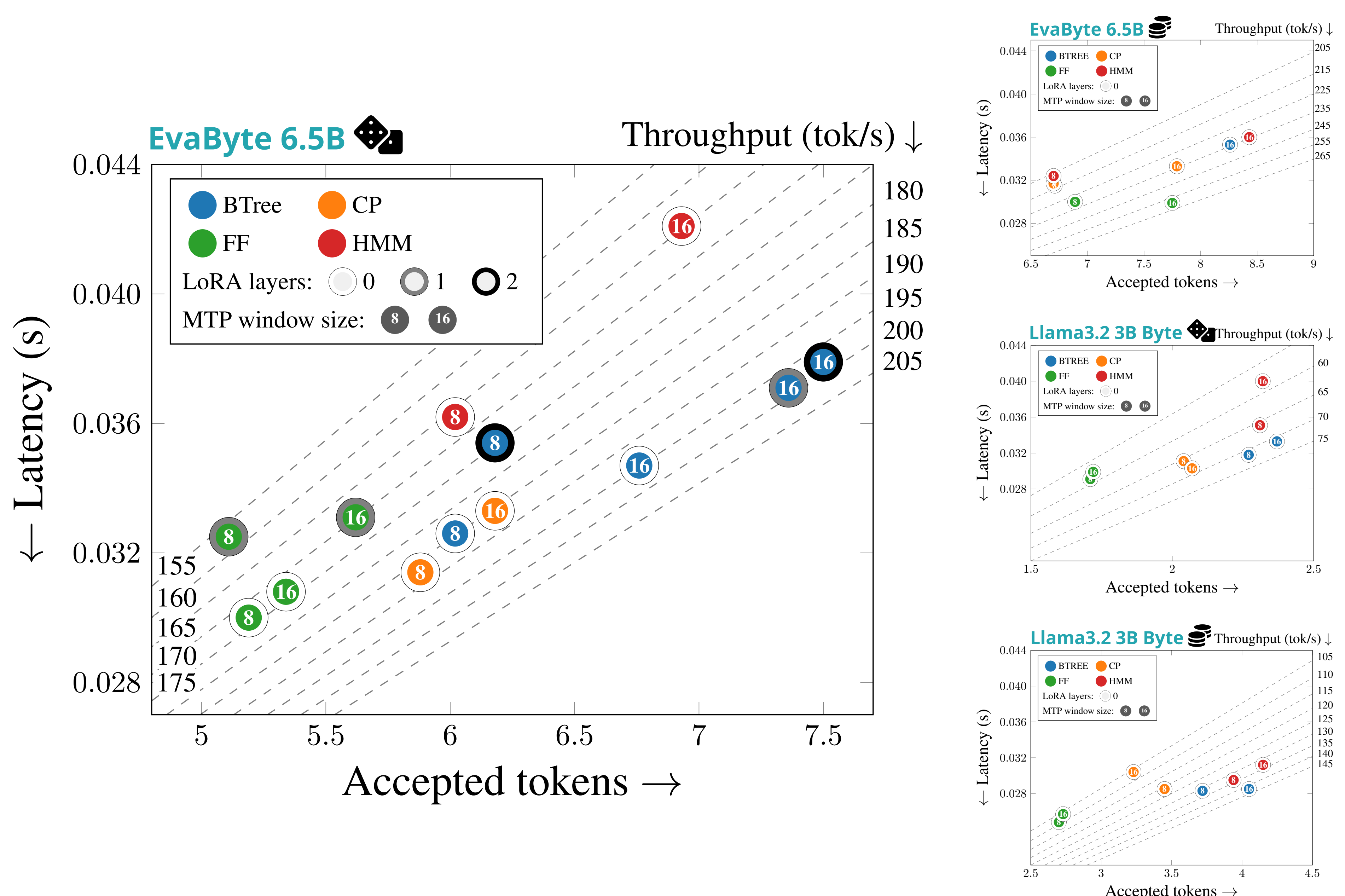
The Expressiveness/Efficiency Trade-off

To increase **throughput** we need to balance **acceptance rate** and **latency**:

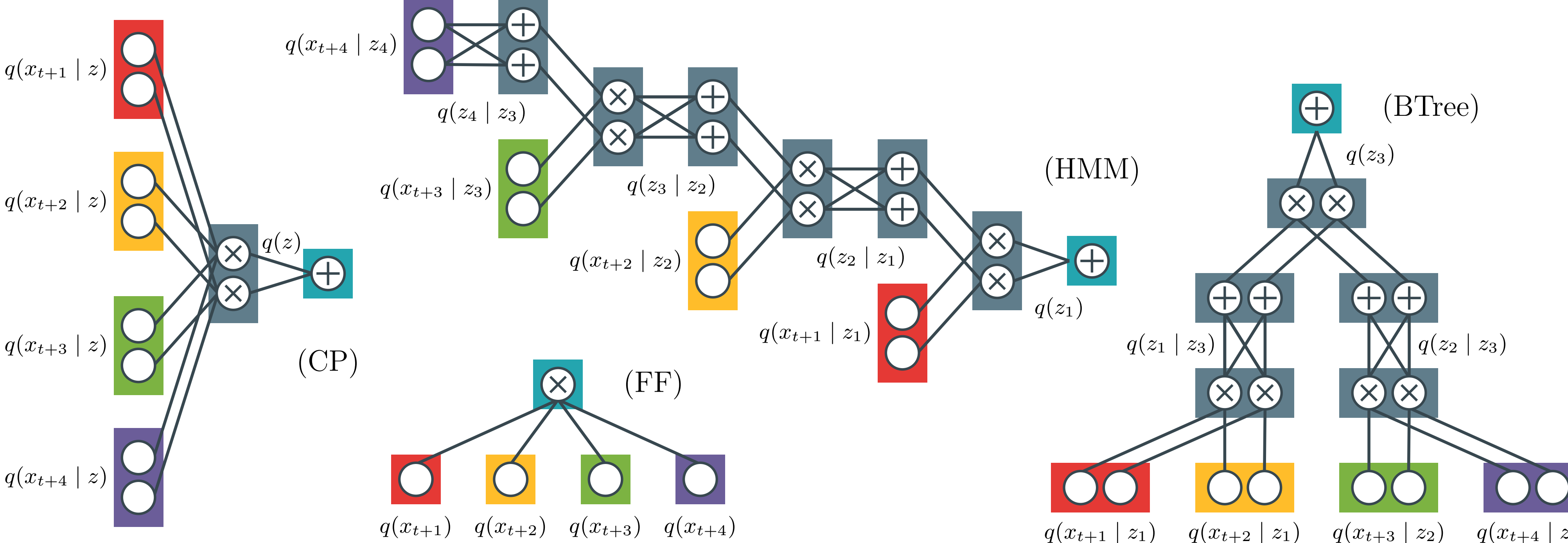


$$\mu_{\text{tok/s}} \approx \frac{\mu_{\text{acc}}}{\mu_{\text{lat}}}$$

TL;DR We speed up a 6.5B byte-level LLM by $5.15\times$ by trading off expressiveness & efficiency with Probabilistic Circuits.



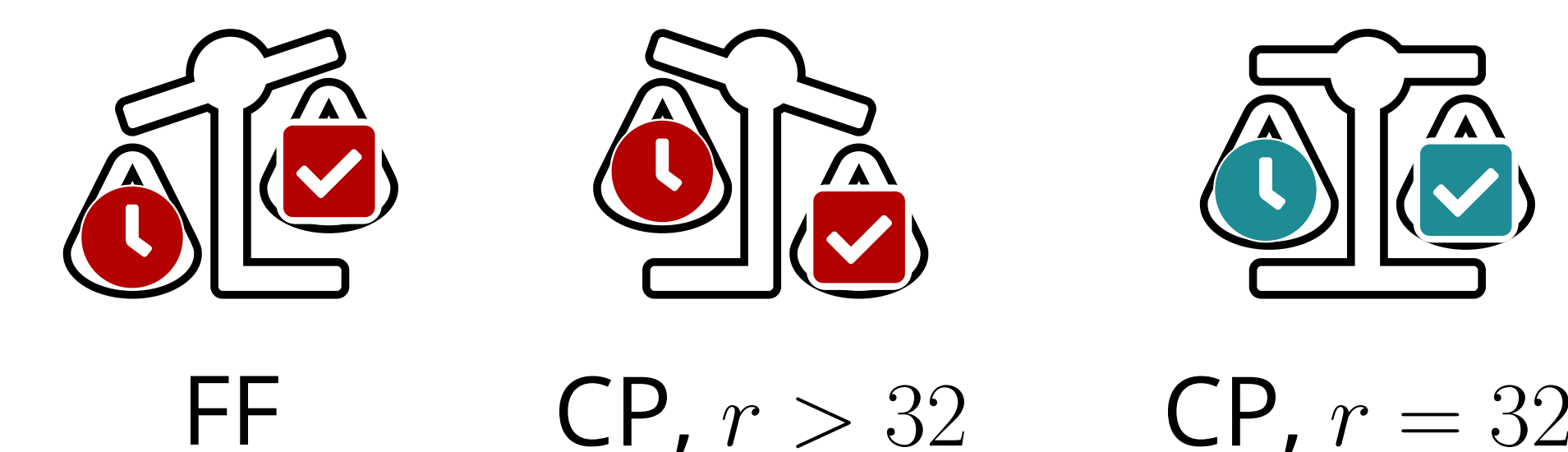
3 Probabilistic Circuits (PCs)



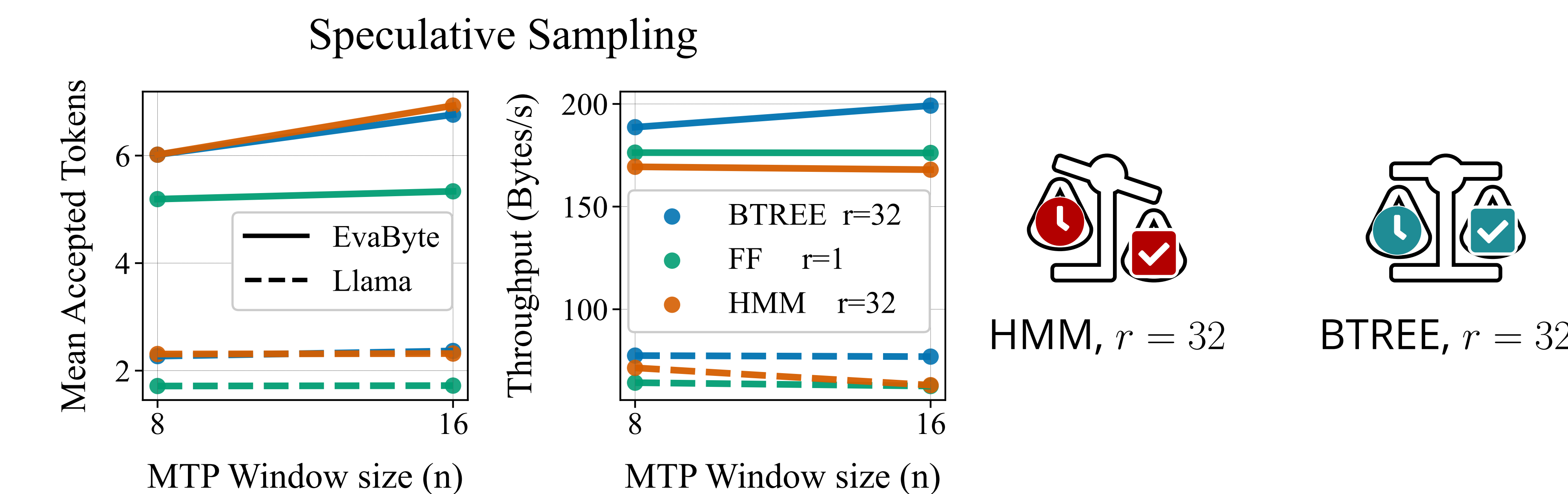
Results (Navigating Expressiveness Axes)

Axis 1 More Mixture Components

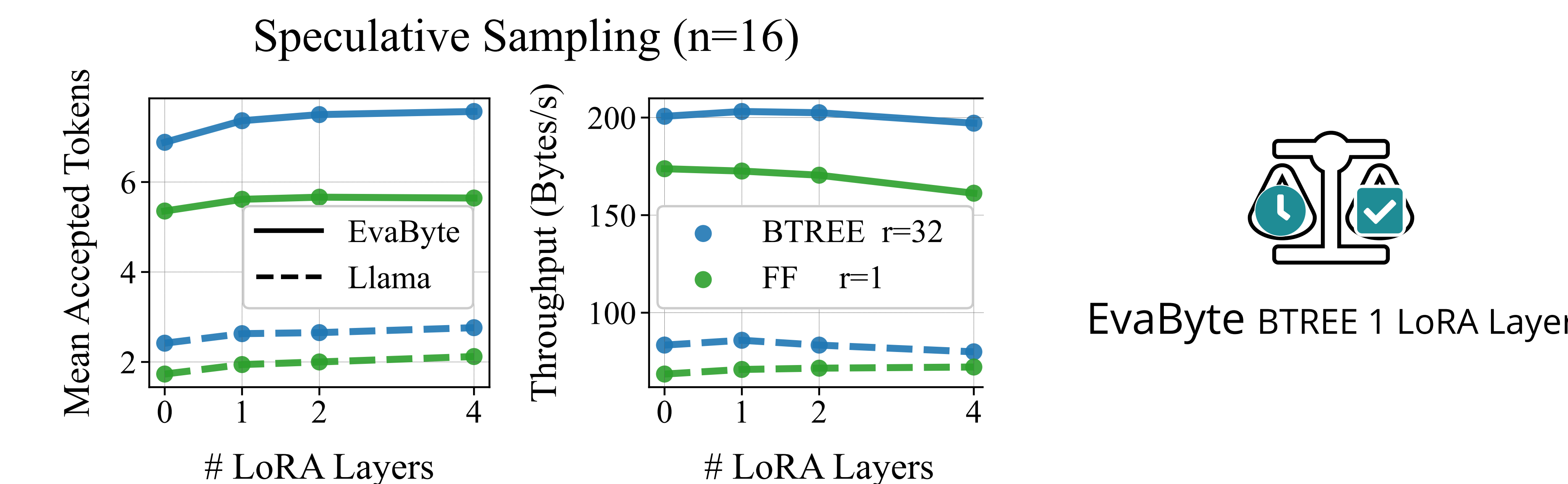
PC	r	$\mu_{\text{lat}} \downarrow$	$\mu_{\text{acc}} \uparrow$	$\mu_{\text{tok/s}} \uparrow$
FF	1	0.0299 ± 0.0003	5.19 ± 0.05	176.3 ± 0.9
	8	0.0314 ± 0.0005	5.67 ± 0.05	184.2 ± 3.9
	16	0.0323 ± 0.0020	5.79 ± 0.08	183.4 ± 13.3
CP	32	0.0314 ± 0.0003	5.88 ± 0.02	191.0 ± 1.6
	64	0.0327 ± 0.0015	5.96 ± 0.04	185.9 ± 7.7
	128	0.0337 ± 0.0001	6.02 ± 0.03	182.3 ± 1.1



Axis 2 More Expressive Circuits



Axis 3 Draft LoRA Layers



Overall Speed-Ups with MTPC

EvaByte 6.5B: $5.15\times$ over Single Byte and $1.17\times$ over FF.
Llama3.2 3B Byte: $2.24\times$ over Single Byte and $1.17\times$ over FF.